

Lecture 7:

Mean-field neural networks

Summary of the class so far: (Mid-semester "sanity check")

2 central questions: { Why optimization is tractable?
Why do NNs generalize?

Two hypotheses

* Tractability via overparametrization: the highly non-convex landscape simplify dramatically as $\# \text{ parameters} \gg \# \text{ samples}$ so that local search methods (e.g., GD/SGD) succeed at finding global minima ("interpolating" solution)

* Implicit regularization of training algorithm: algo itself bias the selection of a solution that generalizes well on new data DESPITE the model being overparametrized and interpolating noisy data.

Lectures 2/3/4: Overparametrization is not incompatible with generalization

Lecture 2: capacity of NNs (Rademacher complexity) depends on norms/spectral properties of the weights and not # parameters.

Lecture 3/4: solution selected by GD (min-norm interpolating soln^o) in linear models can overfit benignly thanks to a self-induced regularization

This motivates the study of gradient-trained NNs that are highly-overparametrized, and as a first approximation $\# \text{ neurons} \rightarrow \infty$

2-layer NN:
$$f(x; \theta) = \frac{1}{\sqrt{M}} \sum_{j=1}^M a_j \sigma(\langle w_j, x \rangle) \quad \|x\|_2 = 1$$

“Standard Xavier initialization” (“used” in pytorch)
 (This is more subtle than this $\rightarrow$ see next week)

$a_j \stackrel{\text{iid}}{\sim} N(0, 1) \quad w_j \stackrel{\text{iid}}{\sim} N(0, \text{Id}_d)$

Lecture 5: As $M \rightarrow \infty$, GD on this NN effectively behave as a linear model

$$f(x; \theta^t) \approx f(x; \theta^0) + \langle \theta^t - \theta^0, \nabla_{\theta} f(x; \theta^0) \rangle \quad (3)$$

CV exponentially fast to global solution (provably trained efficiently)

↳ "Lazy regime" $\|\theta^t - \theta^0\|_2 \ll 1$

(general mechanism that has nothing to do with NNs a priori)

Lecture 6: Linearized NNs are kernel methods which have limited adaptivity (adapt to "smoothness" but not to other low-dimensional structure)

→ "fixed feature" method $\nabla_{\theta} f(x; \theta_0)$
 "fixed representation"

→ suffer from the curse of dimensionality

If θ^t moves away from θ^0 during dynamics, can hope to learn good representation $\nabla_{\theta} f(x; \theta^t)$ and do much better

"feature learning", "representation learning", "learning a good kernel"

(4)

Example: Learning a single neuron

Consider $y_i = \sigma(\langle w_*, x_i \rangle) + \varepsilon_i$ $x_i \sim \text{Unif}(S^{d-1})$
 $\|w_*\|_2 = \sqrt{d}$

Fit with $f(x; \theta) = \frac{1}{\sqrt{M}} \sum_{j=1}^M a_j \sigma(\langle w_j, x \rangle)$

GD on $\hat{R}_n(\theta) = \frac{1}{n} \sum_{i=1}^n (y_i - f(x_i; \theta))^2$

Thm [Montaneri, Zhong, '22, Ghorbani, Mei, Misiakiewicz, Montanari '21]

Fix $\delta > 0$ and take $Md \geq n^{1+\delta}$. Let $\hat{\theta}$ be the solution found in the Lazy learning regime ($\alpha \gtrsim \sqrt{\frac{n^2}{Md}}$)

If $n \leq d^{1+\delta}$, then

$$\liminf_{n, d \rightarrow \infty} R(f(\cdot; \hat{\theta})) \geq \|P_{\mathcal{H}} \sigma\|_{L^2}^2$$

fit at most a d^δ polynomial approximat^o to σ

This is disappointing given that $\sigma(\langle w_*, x \rangle)$ is

(5)

extremely simple and we can fit with one neuron !!!

The problem is w_j^0 don't move during optimization in the lazy regime and we are trying to fit $\sigma(\langle w_*, x \rangle)$ with $\sigma(\langle w_j^0, x \rangle)$

$$\hookrightarrow \text{in high dimensions } \sup_{j \in [M]} \frac{|\langle w_j^0, w_* \rangle|}{\|w_j^0\|_2 \|w_*\|_2} \lesssim \frac{1}{\sqrt{d}}$$

If instead $M=1$ $f(x; \theta) = \sigma(\langle x, w \rangle)$

$$\hat{R}_n(w) = \frac{1}{n} \sum_{i=1}^n (y_i - \sigma(\langle x_i, w \rangle))^2$$

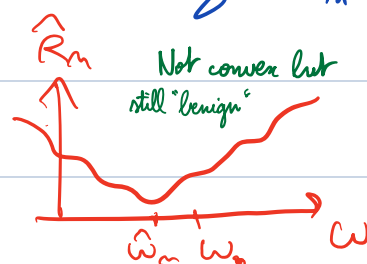
Thm: [Bai, Mei, Montanari, '18] Assume σ bounded and 3 times differentiable bounded derivatives, and $\sigma'(t)$ for all $t \in \mathbb{R}$. Then there exists a constant $C > 0$ such that if

$n \geq C d \log d$, we have

(i) $\hat{R}_n(w)$ has unique local/global minimizer $\hat{w}_n$

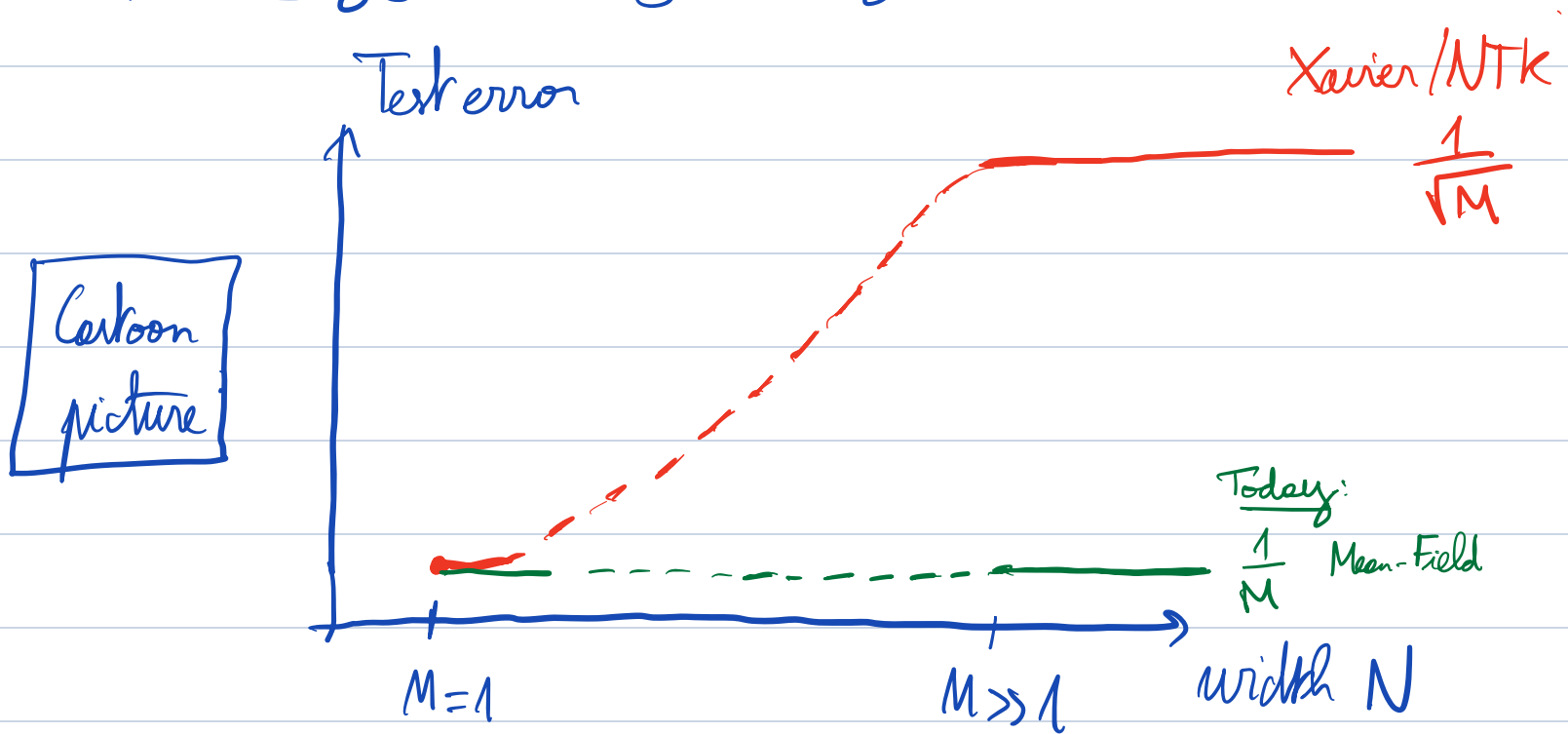
(ii) GD CV to $\hat{w}_n$

(iii) $R(\hat{w}_n) \leq C \sqrt{\frac{d \log n}{n}}$



(6)

Summary: learning a single neuron



Test error degrades as $M \rightarrow \infty$. What is going wrong?

→ This is only a specific infinite-width limit

↳ can choose different scalings and obtain other limits with different behavior.

Today: Mean-Field limit; scaling $\frac{1}{M}$

Mean-Field scaling

$$f(x; \Theta) = \frac{1}{M} \sum_{j=1}^M a_j \sigma(\langle \omega_j, x \rangle)$$

$$\frac{\alpha_M}{\sqrt{M}}$$

$$\alpha_M := \frac{1}{\sqrt{M}}$$

Why is this a natural parametrization?

① Lecture 1: approximation theory of NNs

$$\frac{1}{M} \sum_{j=1}^M a_j \sigma(\langle \omega_j, x \rangle) \xrightarrow{M \rightarrow \infty} \int a(\omega) \sigma(\langle \omega, x \rangle) \mu(d\omega)$$

f_* well approximated with $a(\omega)\mu(d\omega) \rightarrow$ then sample

② Lecture 2: Rademacher complexity $\lesssim n_0 \sqrt{\frac{d}{m}}$ $\frac{\|a\|_1}{M} \leq n_0$
with MF scaling.

$$\left[\text{With NTK scaling } \lesssim n_0 \sqrt{\frac{Md}{m}} \quad \frac{\|a\|_1}{M} \leq n_0 \right]$$

③ Lecture 5: $\frac{\alpha_M}{\sqrt{M}}$ lazy regime $\alpha_M \gtrsim \sqrt{\frac{n^2}{Md}}$

To not be in the lazy regime as $M \rightarrow \infty$ for n/d fixed

we need $\alpha_M \lesssim \frac{1}{\sqrt{M}}$ as $M \rightarrow \infty$

Rmk: Systematic study of infinite width limits of NNs through "ABC" parametrization [Greg Yang, 2021]

- classify scalings that give stable and non-trivial limits as $M \rightarrow \infty$
- MF is special among these limits (for 2-layer) as it maximizes "feature learning"

(Might cover it next week... very useful (easy) computation)

(9)

Stochastic Gradient Descent (SGD)

Let's rewrite $f(x; \Theta) = \frac{1}{M} \sum_{j=1}^M \sigma(x; \theta_j) \quad \theta_j \in \mathbb{R}^D$

$$\Theta = (\theta_1, \dots, \theta_M) \in \mathbb{R}^{MD}$$

e.g.: $\sigma(x; \theta_j) = a_j \tilde{\sigma}(\langle x, w_j \rangle) \quad \theta_j = (a_j, w_j) \in \mathbb{R}^{d+1}$

Consider squared loss: $l(y, \hat{y}) = \frac{1}{2} (y - \hat{y})^2$

Test error: $R_D(\Theta) = \frac{1}{2} \mathbb{E}_{(y,x) \sim D} [(y - f(x; \Theta))^2]$

Train error: $\hat{R}_m(\Theta) = \frac{1}{m} \sum_{i=1}^m \frac{1}{2} (y_i - f(x_i; \Theta))^2$

$$= R_{\hat{D}_m}(\Theta) \quad \hat{D}_m = \frac{1}{m} \sum_{i=1}^m \delta_{y_i, x_i}$$

Simplest algorithm: Gradient Descent (≈ 1850)

$$\Theta^{k+1} = \Theta^k - \eta \nabla \hat{R}_m(\Theta^k)$$

≈ 1950 : we do not need to compute exact gradient

$$\Theta^{k+1} = \Theta^k - \eta \left(\nabla \hat{R}_n(\Theta^k) + z^k \right)$$

$\uparrow$
iid noise with zero mean

$\rightarrow$ noise averages out and algo converges as GD (with suitable step sizes)

SGD (with batch size = 1)

* Initializing: Θ^0 $\Theta_j^0 \sim_{\text{iid}} \mathcal{C}_0$

* At each step: sample (y_k, x_k)

$$\Theta^{k+1} = \Theta^k - \eta M \nabla_{\Theta} l(y_k, f(x_k, \Theta^k))$$

scaling that gives $\Theta_n(t)$ updates

Two cases (1) Finite training samples (y_k, x_k) drawn uniformly at random from $\hat{\mathcal{D}}_n$

"multi-pass SGD"

as $\eta \downarrow 0$ Θ^t follows GF on $\hat{R}_n(\Theta)$

$$\frac{d}{dt} \Theta^t = -M \nabla_{\Theta} \hat{R}_n(\Theta^t)$$

(2) Infinite training samples $(y_k, x_k) \stackrel{\text{i.i.d.}}{\sim} \mathcal{D}$

"one-pass SGD" $\rightarrow$ each sample only used once
 as $\eta \downarrow 0$ $\frac{d}{dt} \Theta_t^t = -M \nabla_{\Theta} R(\Theta^t)$ "population GF"

Both cases can be treated in a unified way $\hat{R}_m = R_{\hat{\mathcal{D}}_m}$ (population dist is $\hat{\mathcal{D}}_m$)
 $\hookrightarrow$ below "one pass SGD"

Dynamics $\cdot \Theta_j^0 \stackrel{\text{i.i.d.}}{\sim} p_0$

$$\cdot \Theta_j^{k+1} = \Theta_j^k + \eta \nabla_{\Theta_j} \sigma(x_k; \Theta_j) \left(y_k - \frac{1}{M} \sum_{\Delta=1}^M \sigma(x_k; \Theta_{\Delta}) \right)$$

Goal: we want to produce a simplified description of the SGD dynamics in this scaling

$$* \underline{\eta \downarrow 0}: \frac{d}{dt} \Theta_t^t = -M \nabla_{\Theta} R(\Theta^t)$$

$$* \underline{M \rightarrow \infty}: f(x; \Theta^t) = f(x; \hat{p}_t^{(M)})$$

$$f(x; p) = \int \sigma(x; \theta) p(\theta) \quad \hat{p}_t^{(M)} = \frac{1}{M} \sum_{j=1}^M \delta_{\Theta_j^t}$$

Empirical weight distribution

Note that $\Theta_j^0 \stackrel{\text{iid}}{\sim} c_0$ at $t=0$ $\hat{c}_0^{(N)} \Rightarrow c_0$

"weak convergence of measures"

It is reasonable to expect

$$\hat{c}_t^{(N)} \Rightarrow c_t \text{ for all } t \geq 0.$$

[i.e., $\forall$ bounded C^∞ fct]
 $\int f d\hat{c}_t^{(N)} \rightarrow \int f dc_t$

where c_t is a deterministic measure following some PDE which we will describe next.

Remark: This is classically referred to as "propagation of chaos":
 as $N \rightarrow \infty$, a system of interacting particles can be well approximated by independent particles interacting with a "mean-field" c_t

Mean-field limit: static properties

Risk: $R_M(\Theta) = \frac{1}{2} \mathbb{E}_{y_g} \left[(y - \beta(x; \Theta))^2 \right]$

$$= \frac{1}{2} \left[R_{\#} + \frac{2}{M} \sum_{j=1}^M V(\theta_j) + \frac{1}{M^2} \sum_{i,j=1}^M U(\theta_i, \theta_j) \right]$$

where

$$R_{\#} = \mathbb{E}[y^2] \quad V(\theta) = -\mathbb{E}_{y_g} [y \sigma_{\#}(x; \theta)]$$

$$U(\theta_1, \theta_2) = \mathbb{E}_x [\sigma_{\#}(x; \theta_1) \sigma_{\#}(x; \theta_2)]$$

Interpretation: $R_M(\Theta)$ is the energy of a system of M particles $\theta_j \in \mathbb{R}^D$, interacting with an external potential $V(\theta)$ and via a pairwise potential $U(\theta_1, \theta_2)$

In particular: U is P.S.D. i.e. for any bounded and compactly supported function h , we have

$$\int h(\theta_1) U(\theta_1, \theta_2) h(\theta_2) d\theta_1 d\theta_2 \geq 0$$

"repulsive interaction".

Gradient w.r.t one particle:

(14)

$$M \nabla_{\theta_j} R_M(\Theta) = \nabla_{\theta_j} \left(V(\theta_j) + \frac{1}{M} \sum_{i=1}^M U(\theta_j, \theta_i) \right) \\ =: \nabla_{\theta_j} \Psi(\theta_j; \Theta)$$

As M is large, it is natural to replace $\theta_1, \dots, \theta_M$ by a density $\rho \in \mathcal{P}(\mathbb{R}^D)$

$$R(\rho) = \frac{1}{2} R_{\#} + \int V(\theta) \rho(d\theta) + \int U(\theta_1, \theta_2) \rho(d\theta_1) \rho(d\theta_2)$$

$$\text{And } \Psi(\theta; \rho) = V(\theta) + \int U(\theta, \theta') \rho(d\theta')$$

$$= \frac{\delta R(\rho)}{\delta \rho(\theta)}$$

("additional energy of adding a single particle at $\theta \in \mathbb{R}^D$ ")

Remark: $R(\hat{\rho}^{(M)}) = R_M(\Theta)$

$$\frac{d}{dt} \theta_j^t = -\nabla_{\theta} \Psi(\theta_j; \hat{\rho}_t^{(M)})$$

(this is a completely equivalent description of GF on M neurons)

Hence, heuristically, $\hat{\rho}_t^{(M)} \Rightarrow \rho_t$

Each particle

$$\theta^t \sim \rho_t \text{ satisfy } \frac{d}{dt} \theta^t = -\nabla_{\theta} \Psi(\theta^t; \rho_t)$$

"consistency equation"

PDE characterization

We can describe evolution of p_t in terms of a PDE

Below is a heuristic derivation: assume $p_t(d\theta) = \overset{\text{has a density}}{p_t(\theta)} d\theta$
^{for simplicity}

Consider a differentiable test function $g: \mathbb{R}^D \rightarrow \mathbb{R}$

$$\begin{aligned} \frac{d}{dt} \mathbb{E}[g(\theta^t)] &= \frac{d}{dt} \int g(\theta) p_t(\theta) d\theta \\ &= \int g(\theta) \partial_t p_t(\theta) d\theta \end{aligned} \quad \begin{array}{l} \downarrow \text{differentiating} \\ \text{under integral} \end{array}$$

On the other hand, exchanging derivative / expectation

$$\begin{aligned} \frac{d}{dt} \mathbb{E}[g(\theta^t)] &= \mathbb{E}\left[\frac{d}{dt} g(\theta^t)\right] = \mathbb{E}\left[\langle \nabla_{\theta} g(\theta^t), \frac{d}{dt} \theta^t \rangle\right] \\ &= - \mathbb{E}\left[\langle \nabla_{\theta} g(\theta^t), \nabla_{\theta} \psi(\theta^t; p_t) \rangle\right] \end{aligned}$$

$$\begin{aligned} &= - \int \langle \nabla_{\theta} g(\theta), \nabla_{\theta} \psi(\theta; p_t) \rangle p_t(\theta) d\theta \\ &= \int g(\theta) \underbrace{\nabla \cdot}_{\uparrow} [p_t(\theta) \nabla_{\theta} \psi(\theta; p_t)] d\theta \end{aligned}$$

Integrate by part $\downarrow$

$$\text{divergence } \nabla_{\theta} \cdot v(\theta) = \sum_{i=1}^D \frac{\partial}{\partial \theta_i} v_i(\theta) \quad (16)$$

Hence putting equalities together:

$$\int g(\theta) \left\{ \partial_t c_t(\theta) - \nabla \cdot [c_t(\theta) \nabla_{\theta} \Psi(\theta; c_t)] \right\} d\theta = 0$$

"Weak form" of the following PDE

$$\partial_t c_t = \nabla \cdot [c_t \nabla_{\theta} \Psi(\theta; c_t)]$$

This is often referred to as "Distributional Dynamics"

This is known as a McKean-Vlasov type PDE.

PDE independently proven in 4 concurrent work
[CB'18, MMN'18, RVE'18, SS'18]

If initialize a particle $\theta_j^0 \sim c_0$ and

$$\frac{d}{dt} \theta_j^t = -\nabla_{\theta} \Psi(\theta_j^t; c_t) \quad \text{then } \theta_j^t \sim c_t \text{ for all } t \geq 0.$$

Approximation guarantees

How well $f(x; \rho_t)$ tracks $f(x; \Theta^k)$?

- where $\cdot$
- Θ^k solut^o of SGD with M neurons
 η step size
 - ρ_t solut^o of above PDE.

Assumptions: (i) $|\sigma|, |y| \leq C$, $\nabla_{\theta} \sigma(x; \theta)$ C-subGaussian
(ii) $V(\theta)$, $U(\theta_1, \theta_2)$ have bounded Lipschitz gradient
(weaker assumpt^o in the paper)

Thm [Mei, Misakiewicz, Montanari, '19]
There exists a constant K that only depends on assumptions
s.t. for all $T \geq 0$, with probability $\geq 1 - e^{-\delta}$

$$\sup_{k \in [0, \frac{T}{\eta}] \cap \mathbb{N}} |R_M(\Theta^k) - R(\rho_{k\eta})| \leq \frac{Ke^{KT}}{\sqrt{M}} (\sqrt{\log M} + 3) + Ke^{KT} (\sqrt{D + \log M} + 3) \sqrt{\eta}$$

error due to $M \rightarrow \infty$ approx

error due to $\eta \rightarrow 0$ approx.

Remark: • $\hat{\ell}_k^{(M)} \approx \ell_k \eta$ $t = k\eta$

- For T fix, $M \gtrsim 1$ and $\eta \lesssim \frac{1}{D}$ for MF to be a good approximation.

→ In particular M does not need to depend on D but only on properties of the target fct

→ similar to approximation theory (e.g. Borron's result)

- e^{kT} is bad but in some sense unavoidable

$$\hookrightarrow T = O(1) \quad \frac{T}{\eta} = O(D) \text{ SGD steps}$$

→ good approx limited to this linear scaling

→ will talk more about it in following lectures

Wasserstein gradient flow description

The PDE has an interesting structure: it is a gradient flow in the space of distributions.

It is worth fleshing out this relation as it connects MF limit to a rich field of mathematics: optimal transport.

What is a gradient flow?

Consider:

* (M, ρ) a metric space ($\rho = \text{distance}$)

* $F: M \rightarrow \mathbb{R}$

Trajectory: $z(0), z(\varepsilon), z(2\varepsilon), \dots, z(k\varepsilon), \dots$

$$z(k\varepsilon + \varepsilon) = \underset{z \in M}{\operatorname{argmin}} \left\{ F(z) + \frac{1}{2\varepsilon} \rho(z, z(k\varepsilon)) \right\}$$

As $\varepsilon \searrow 0$, this defines a continuous trajectory $z(t)$

"GF on F in (M, ρ) "

The minimizing "movements" in (M, ρ) do depend on $F(z)$

E.g. Euclidean space $\frac{dz(t)}{dt} = -\nabla F(z(t))$ is indeed the limit of the above

Previous PDE is a Gradient flow on $R(\rho)$ in the metric space $(\mathcal{P}_2(\mathbb{R}^D), W_2)$

* $\mathcal{P}_2(\mathbb{R}^D)$: space of probe ρ on $\mathbb{R}^D$ with $\int \|\theta\|_2^2 \rho(d\theta) < \infty$

* W_2 : Wasserstein distance

$$W_2(\mu, \nu) = \inf_{\gamma \in \Gamma(\mu, \nu)} \left(\int \|\alpha - y\|_2^2 \gamma(dx, dy) \right)^{\frac{1}{2}}$$

where $\Gamma(\mu, \nu)$ is the set of all couplings of μ and ν
i.e. joint probe on $\mathbb{R}^D \times \mathbb{D}$ with marginals μ and ν

$$\rho_{t+\varepsilon} \approx \operatorname{argmin}_{\rho \in \mathcal{P}(\mathbb{R}^D)} \left\{ R(\rho) + \frac{1}{2\varepsilon} W_2(\rho, \rho_t)^2 \right\}$$

In particular $\frac{dR(\rho_t)}{dt} = - \int \|\nabla \varphi(\theta; \rho_t)\|_2^2 \rho_t(d\theta) \leq 0$

Global convergence

$R(p_t)$ is monotonically decreasing. Can we show global convergence? (And have some success as lazy regime)

$$\lim_{t \rightarrow \infty} R(p_t) = \inf_p R(p)$$

[Here we consider one-pass SGD hence global CV is directly on test error, i.e., directly implies that we generalize well]

→ not quite: showing global CV of this Wasserstein flow is hard and there are many subtle difficulties.

Below I describe some basic properties.

Static properties:

Minimizing $R_M(\odot)$ is not much different than minimizing $R(p)$

Lemma: Assume $\exists \varepsilon, K > 0$ such that

$$R(p) \leq \inf_e R(p) + \varepsilon \Rightarrow \int U(\theta, \theta) p(d\theta) \leq K$$

Then:

$$\left| \inf_{\Theta \in \mathbb{R}^{MD}} R_M(\Theta) - \inf_e R(p) \right| \leq \frac{K}{M}$$

Proof: Take p_* s.t. $R(p_*) \leq \inf_e R(p) + \delta$ $\delta \leq \varepsilon$
(in particular $\int U(\theta, \theta) p_*(d\theta) \leq K$ by assumption)

Consider $(\theta_j)_{j \in [M]} \sim \text{iid } p_*$, then simple calculation yields

$$\begin{aligned} \mathbb{E}_{p_*} [R_M(\Theta)] - R(p_*) &= \frac{1}{M} \left\{ \int U(\theta, \theta) p_*(d\theta) - \underbrace{\int U(\theta_1, \theta_2) p_*(d\theta_1) p_*(d\theta_2)}_{\geq 0 \text{ by PSD}} \right\} \\ &\leq \frac{1}{M} \int U(\theta, \theta) p_*(d\theta) \leq \frac{K}{M} \end{aligned}$$

Hence $\inf_{\Theta} R_M(\Theta) \leq \inf_e R(p) + \frac{K}{M} + \delta$ for any δ $\square$

We further have the following characterization of minimizer

Prop: Assume U, V are C^0 and bounded from below
 Assume $\exists \varepsilon, \gamma, K > 0$, s.t.

$$R(\varrho) \leq \inf_{\varrho} R(\varrho) + \varepsilon \Rightarrow \int \|\Theta\|^\gamma \varrho(d\Theta) \leq K$$

Then (1) $\inf_{\varrho} R(\varrho)$ is achieved at some ϱ_* .

(2) ϱ_* is a minimizer iff
 $\text{supp}(\varrho_*) \subseteq \arg\min_{\Theta} \Psi(\Theta; \varrho_*)$

(Proof is by doing a perturbation of ϱ_*)

Corollary: Suppose ϱ_* satisfy (i) $\text{supp}(\varrho_*) = \mathbb{R}^D$
 and (ii) $\Psi(\Theta; \varrho_*)$ is a constant
 Then ϱ_* is a minimizer.

Proof: (This is direct from previous proposition (2))

Dynamical properties:

ρ is a fixed point of the PDE if and only if

$$(*) \quad \text{supp}(\rho) \subseteq \{\theta : \nabla_{\theta} \mathcal{U}(\theta; \rho) = 0\}$$

(simply comes from $\frac{d}{dt} R(\rho_t) = - \int \|\nabla_{\theta} \Psi\|_2^2 \rho_t(d\theta)$)

There will be ∞ number of "bad" stationary points
(e.g. any stationary pt of landscape with finite # of neurons will be a stationary of the PDE with empirical measure $\hat{\rho}^{(n)}$)

It is hard & subtle to argue that ρ_t will not get trapped in these stationary points.

Here is a simple result:

Thm: Assume (a) $\rho_t \rightarrow \rho_{\infty} \in \mathcal{P}_2(\mathbb{R}^D)$

(b) $\text{supp}(\rho_{\infty}) = \mathbb{R}^D$

(+ U, V differentiable with bounded gradient)

Then ρ_{∞} is a minimizer, that is

$$R(\rho_t) \rightarrow \min_{\rho} R(\rho)$$

Proof: ρ_∞ stationary point, i.e.
 $\mathbb{R}^D = \text{supp}(\rho_\infty) \subseteq \{\theta: \nabla \Psi(\theta; \rho_\infty) = 0\}$
 hence $\Psi(\theta; \rho_\infty)$ cste and we conclude using
 the corollary above $\square$

Assumption $\text{supp}(\rho_\infty) = \mathbb{R}^D$ is very restrictive

[Chizat, Bach, '18] [Wojtowysch, '22]

[Nguyen, Pham, '20]

Always need to assume $\left[\begin{array}{l} \rho_t \rightarrow \rho_\infty \quad (*) \\ \rho_\infty \text{ has some abstract properties} \end{array} \right.$

$\left[\begin{array}{l} (*) \text{ this is not a benign condition as } \rho_t \text{ might oscillate while} \\ R(\rho_t) \searrow R_* \end{array} \right]$

Nonetheless: these proofs are instructive and show that it is
 hard to trap a distribution in a non-optimal stationary point
 (lockability via overparameterization)

Summary: • No quantitative general CV guarantees
 • Even qualitative guarantees are under abstract assumptions

Noisy SGD

26

Consider the following variant of SGD
 → add $\frac{\lambda}{2} \|\Theta\|_2^2$ regularization to the risk

→ add at each step Gaussian noise to the gradient

At each step: $\Theta_j^{k+1} = (1 - \lambda\eta)\Theta_j^k + \eta \nabla_{\Theta} \phi(x_k; \Theta_j^k) (y_k - \phi(x_k; \Theta_j^k))$
 $+ \sqrt{\frac{\eta}{\beta}} g_j^k$
 $\rightarrow g_j^k \stackrel{iid}{\sim} N(0, I_D)$

Risk: $R_{\lambda}(\rho) = R(\rho) + \frac{\lambda}{2} \int \|\Theta\|_2^2 \rho(d\Theta)$

$$\Psi_{\lambda}(\theta; \rho) = \Psi(\theta; \rho) + \frac{\lambda}{2} \|\theta\|_2^2$$

Then as $M \rightarrow \infty$ $\eta \downarrow 0$

The resulting dynamics

$$\partial_t \rho_t = \nabla_{\Theta} \cdot [\rho_t \nabla_{\Theta} \Psi_{\lambda}(\theta; \rho_t)] + \frac{1}{\beta} \Delta \rho_t$$

Additional diffusion term

"inverse temperature"

$$\Delta \cdot f(\theta) = \sum_{i=1}^D \frac{\partial^2}{\partial \theta_i^2} f(\theta)$$

This is also a Wasserstein gradient flow but now on the free energy

$$F_{\lambda, \beta}(e) = R_{\lambda}(e) - \frac{1}{\beta} S(e)$$

with $S(e) = - \int e(\theta) \log e(\theta) d\theta$ the entropy of e

The interesting property when adding entropic regularization is that the Free energy is strongly convex and there is only one stationary point which is the global minimizer of $F_{\lambda, \beta}(e)$

Lemma: For any $\lambda \geq 0$ and $\beta > 0$, $F_{\lambda, \beta}(e)$ has at most one fixed point $F_{\lambda, \beta}(e_*) < \infty$.

Furthermore e_* has density satisfying self-consistent Boltzmann equation

$$e_*(\theta) = \frac{1}{Z(\lambda, \beta)} \exp(-\beta \psi_{\lambda}(\theta; e_*))$$

Proof: This follows from showing that the dissipat^o equation

$$\frac{d F_{\lambda, \beta}(e_t)}{dt} = - \int_{\mathbb{R}^D} \left\| \nabla_{\theta} \left[\psi_{\lambda}(\theta; e_t) + \frac{1}{\beta} \log e_t(\theta) \right] \right\|_{2, e_t(\theta)}^2 d\theta$$

It is $= 0$ iff e^* satisfy equation above

In particular this means that

$$\lim_{t \rightarrow \infty} F_{\lambda, \beta}(e_t) = \min_e F_{\lambda, \beta}(e)$$

If β taken sufficiently small $\min_e F_{\lambda, \beta}(e) \approx \min_e R_{\lambda}(e)$

[We need $\beta = \Omega(D)$]

We get: $\lim_{t \rightarrow \infty} R(e_t) \leq \inf_e R(e) + \mathcal{O}\left(\frac{D}{\beta}\right)$

Great !!! Global CV as long as $\beta < \infty$
but what about convergence time?

(29)

Standard approach: prove Log-Sobolev inequality
for this setting

→ If LSI with LS constant α then

$$F_{\lambda, B}(p_t) - F_{\lambda, B}(p_*) \leq e^{-C \frac{\alpha}{B} t} (F_{\lambda, B}(p_0) - F_{\lambda, B}(p_*))$$

However: best we can show $\alpha \geq \lambda B e^{-CB}$

hence if $B = \Theta(D)$ $\alpha \gtrsim e^{-\Theta(D)}$

and require $T = e^{\Theta(D)}$ to converge

→ this cannot be improved in general (i.e., there exist settings where we will need e^D time to CV).

Returning to single neuron learning example

(30)

We return to the example of learning a single neuron which we started with

This illustrates how the PDE can be simplified in settings with symmetries.

assume $\sigma'(t) > 0 \quad \forall t \in \mathbb{R}$

Consider: $x_i \sim N(0, I_d)$ $y_i = \sigma(\langle w, x_i \rangle) + \varepsilon_i$

Model: $f(x; \theta) = \int a \sigma(\langle w, x \rangle) \rho(da, dw)$

Training: "uninformative initialization"

$$\rho_0(da, dw) = \underbrace{P_A}_{\text{initialized } a} \otimes \underbrace{N(0, \frac{\sigma^2}{d} I_d)}_{\text{initialized } w}$$

$$\partial_t \rho_t = \nabla \cdot [\rho_t \nabla_{\theta} \Psi(\theta; \rho_t)]$$

(31)

Exploiting symmetries: data distribution (y, x) remain invariant under a rotation that leaves w_* unchanged

Hence: if $R \in \mathbb{R}^{d \times d}$ s.t. $R w_* = w_*$

then $R_{\#} p_t$ is a solution of the PDE with initialization $R_{\#} p_0$

Here $R_{\#} p_0 = p_0 \rightarrow$ by uniqueness of the solution

we must have $R_{\#} p_t = p_t$ for all $t \geq 0$.

Hence p_t is invariant under rotations that leave $R w_* = w_*$

i.e. only depend on $(a, \underbrace{\langle w_*, w \rangle}_s, \underbrace{\|P_{\perp} w\|_2}_r)$

$$p_t(da, dw) = \bar{p}_t(da, ds, dr) \otimes \text{Unif}(S^{d-2})$$

Denote $z = (a, s, r)$

$$\Psi(z; \bar{p}_t) = \Psi(\theta; p_t) = V(\theta) + \int U(\theta, \theta') p_t(d\theta')$$

$$V(\theta) = -\mathbb{E}[y a \sigma(\langle w, x \rangle)] \quad U(\theta, \theta') = \mathbb{E}[a a' \sigma(\langle w, x \rangle) \sigma(\langle w', x \rangle)]$$

$$V(\theta) = -a \mathbb{E}[\sigma(\langle \omega, x \rangle) \sigma(\langle \omega, x \rangle)]$$

$$= -a \mathbb{E}_{\substack{G_1, G_2 \\ \sim N(0,1)}} [\sigma(G_1) \sigma(s G_1 + n G_2)]$$

$$U(\theta, \theta') = a a' \mathbb{E}_{\substack{G_1, G_2, G_3 \\ \sim N(0,1)}} \left[\sigma(s G_1 + n G_2) \sigma(s' G_1 + \frac{\langle \omega_{\perp}, \omega'_{\perp} \rangle}{n} G_2 + \sqrt{n'^2 - \frac{\langle \omega_{\perp}, \omega'_{\perp} \rangle^2}{n^2}} G_3 \right]$$

When taking integral over $\omega'_{\perp} = n' u$ $u \sim \text{Unif}(S^{d-2})$
this indeed only depend on (a, s, n) .

Hence the dynamics only depends on 3 parameters instead of $d+1$
Reduced form $\bar{p}_t \in P(\mathbb{R}^3)$

$$\partial_t \bar{p}_t = \partial_a (\bar{p}_t \partial_a \Psi(g_i \bar{p}_t)) + \partial_s (\bar{p}_t \partial_s \Psi(g_i \bar{p}_t)) + \frac{1}{n} \partial_n (n \bar{p}_t \partial_n \Psi(g_i \bar{p}_t))$$

Initialization translates into $\bar{p}_0 = \underbrace{P_1}_{a^0} \otimes \underbrace{N(0, \frac{\gamma^2}{d})}_{s^0} \otimes \underbrace{Q_{d-1, \gamma}}_{n^0}$

where $Q_{k, \gamma} = \frac{\gamma}{\sqrt{k+1}} \sqrt{\chi^2_{(k)}}$

The above PDE can be solved numerically efficiently

$$\lim_{t \rightarrow \infty} R(p_t) = 0$$

excess test error

For $\varepsilon > 0$, $\exists T = T(\varepsilon)$ indep of d s.t. $R(p_t) \leq \varepsilon$

Then taking $M \gtrsim \frac{e^{KT(\varepsilon)}}{\varepsilon^2} = \tilde{O}_d(1)$

$$\eta \lesssim \frac{e^{-KT(\varepsilon)}}{\varepsilon^2 d} = \tilde{O}_d\left(\frac{1}{d}\right)$$

We get $R_M\left(\tilde{W}^{\frac{I}{\eta}}\right) \leq R(p_T) + \varepsilon \leq 2\varepsilon$

Total # samples = $\tilde{O}_d(d)$ # neurons = $\tilde{O}_d(1)$

(Success?)

Conclusion

Mean-Field scaling: $\ell_t \neq \ell_0$ "feature learning"

From single neuron example: can indeed adapt to low-dim structure and vastly outperform NTK.

Goal of $M \rightarrow \infty$ is to simplify analysis of SGD.

Did we succeed?

Pros:

- Compact description of dynamics in terms of PDE
- When setting has symmetries, can exploit them to reduce the effective dimension of the parameters

Cons:

- These PDE are hard to analyze in general
- Hard to prove qualitative/quantitative CV guarantees
- only good approx for $\Theta(D)$ steps of SGD

Ongoing line of work!!!